

# 兵棋推演的智能决策技术与挑战

尹奇跃<sup>1,2</sup> 赵美静<sup>1</sup> 倪晚成<sup>1,2</sup> 张俊格<sup>1,2</sup> 黄凯奇<sup>1,2</sup>

**摘要** 近年来,以人机对抗为途径的智能决策技术取得了飞速发展,人工智能(Artificial intelligence, AI)技术 AlphaGo、AlphaStar 等分别在围棋、星际争霸等游戏环境中战胜了顶尖人类选手. 兵棋推演作为一种人机对抗策略验证环境,由于其非对称环境决策、更接近真实环境的随机性与高风险决策等特点,受到智能决策技术研究者的广泛关注. 通过梳理兵棋推演与目前主流人机对抗环境(如围棋、德州扑克、星际争霸等)的区别,阐述了兵棋推演智能决策技术的发展现状,分析了当前主流技术的局限与瓶颈,对兵棋推演中的智能决策技术研究进行了思考,期望能对兵棋推演相关问题中的智能决策技术研究带来启发.

**关键词** 兵棋推演, 人机对抗, 智能决策技术, 博弈学习

**引用格式** 尹奇跃, 赵美静, 倪晚成, 张俊格, 黄凯奇. 兵棋推演的智能决策技术与挑战. 自动化学报, 2023, 49(5): 913–928

**DOI** 10.16383/j.aas.c210547

## Intelligent Decision Making Technology and Challenge of Wargame

YIN Qi-Yue<sup>1,2</sup> ZHAO Mei-Jing<sup>1</sup> NI Wan-Cheng<sup>1,2</sup> ZHANG Jun-Ge<sup>1,2</sup> HUANG Kai-Qi<sup>1,2</sup>

**Abstract** In recent years, decision-making intelligence based on human-machine confrontation has achieved rapid development. For example, artificial intelligence (AI) technology such as AlphaGo and AlphaStar have defeated top human players in games Go and StarCraft, respectively. Nowadays, wargame, as a new verification environment for human-machine confrontation, attracts more and more researchers due to new challenges being raised, i.e., asymmetric environmental decision-making and randomness with high-risk decision-making. In this paper, we will sort out the differences between wargame and the current mainstream human-machine confrontation environments such as Go, Poker and StarCraft. Then, we explain the development status of wargame intelligent technology, and analyze the limitations of current mainstream technologies. Finally, we present our thoughts about future development of technologies for wargame, hoping to inspire researchers for through study on wargame.

**Key words** Wargame, human-machine confrontation, intelligent decision making technology, game learning

**Citation** Yin Qi-Yue, Zhao Mei-Jing, Ni Wan-Cheng, Zhang Jun-Ge, Huang Kai-Qi. Intelligent decision making technology and challenge of wargame. *Acta Automatica Sinica*, 2023, 49(5): 913–928

作为人工智能(Artificial intelligence, AI)技术的试金石,针对人机对抗的研究,近年来获得了举世瞩目的进展.随着 Deep Blue<sup>[1]</sup>、AlphaGo<sup>[2]</sup>、Libratus<sup>[3]</sup>、AlphaStar<sup>[4]</sup>等智能体分别在国际象棋、围棋、二人无限注德州扑克以及星际争霸中,战胜顶尖职业人类选手,其背后的智能决策技术获得了广泛的关注,也代表了智能决策技术在中等复杂度完美信息博弈<sup>[1]</sup>、高复杂度完美信息博弈<sup>[2]</sup>再到高复杂度不完美信息博弈中的技术突破<sup>[3]</sup>.

国际象棋、围棋代表了完美信息博弈,其状态空间复杂度由  $10^{47}$  增至  $10^{360}$ ,后者更是被誉为人工智能技术的阿波罗<sup>[2]</sup>.相比于上述两种博弈环境,二人无限注德州扑克尽管状态空间复杂度仅有  $10^{160}$ ,但其为不完美信息博弈,相比于国际象棋与围棋信息集大小仅为 1,信息集平均大小达到  $10^3$ <sup>[3]</sup>.而星际争霸作为高复杂度不完美信息博弈的代表,因相比于上述游戏的即时制、长时决策等特性<sup>[4–5]</sup>,对智能决策技术提出了更高的要求.

星际争霸突破之后,研究人员迫切需要新的人机对抗环境实现智能技术的前沿探索.兵棋推演是一款经典策略游戏<sup>[6–8]</sup>,也被称为战争游戏,作为一种人机对抗策略验证环境,由于其具有不对称环境决策、更接近真实环境的随机性与高风险决策等特点,受到智能决策技术学者的广泛关注.近年来,学者们投入了大量的精力进行兵棋推演智能体研发和兵棋推演子问题求解,试图解决兵棋推演的人机对抗挑战<sup>[9–14]</sup>.

收稿日期 2021-06-17 录用日期 2021-09-17

Manuscript received June 17, 2021; accepted September 17, 2021

国家自然科学基金(61906197)资助

Supported by National Natural Science Youth Foundation of China (61906197)

本文责任编辑 黄华

Recommended by Associate Editor HUANG Hua

1. 中国科学院自动化研究所 北京 100190 2. 中国科学院大学 北京 100049

1. Institute of Automation, Chinese Academy of Sciences, Beijing 100190 2. University of Chinese Academy of Sciences, Beijing 100049

兵棋推演一直以来都是战争研究和训练的手段. 自 1811 年冯·莱斯维茨在前人研究成果基础上发明了现代兵棋以来, 兵棋推演迅速流行起来, 并逐渐演化出两个分支, 即“严格式”和“自由式”兵棋推演. 时至今日, 兵棋推演在实战化训练、指挥官培训等任务中扮演着越来越重要的角色<sup>[15-17]</sup>. 胡晓峰等<sup>[6]</sup>全面综述了兵棋推演的基本要素(参演人员、兵棋系统模拟的战场环境和作战部队、导演部及导调机构), 指出“兵棋推演的难点在于模拟人的智能行为”, 进而得出“兵棋推演需要突破作战态势智能认知瓶颈”, 最后给出了如何实现态势理解与自主决策可能的路径. 可以看出, 当前兵棋推演的研究集中在如何更真实模拟对抗以及如何实现态势分析与推理, 进而辅助人类指挥决策.

与目前兵棋推演关注的重点不同, 本文关注兵棋推演的智能体研究, 从人机对抗这一角度出发, 针对通用性的智能决策技术与挑战展开论述. 基于该角度, 将展示兵棋推演这一对人机对抗智能决策技术带来新挑战的研究环境, 以吸引更多的研究者解决兵棋推演人机对抗挑战, 服务于兵棋推演应用(如为指挥员提供多样化完整策略)的同时, 促进智能决策技术的新突破. 另外需要阐明的是, 本文中的兵棋推演, 如非特别阐述, 在不引起歧义的前提下统一指两方计算机兵棋推演(红、蓝两方).

本文结构如下: 第 1 节介绍兵棋智能决策问题的挑战; 第 2 节详细阐述兵棋智能决策技术研究现状, 包括智能体研发技术与框架、智能体评估评测及平台; 第 3 节总结兵棋智能决策技术的挑战; 第 4 节对兵棋智能决策技术进行展望; 第 5 节作为结束语, 简要总结全文.

1 兵棋智能决策问题的挑战

本节首先简要介绍兵棋推演问题以及与手工兵棋的比较. 在此基础上, 以人机对抗发展脉络为主线, 以兵棋推演中的智能体研究为核心, 介绍兵棋推演与其他主流策略游戏的通用挑战. 然后, 重点阐述兵棋推演的独特挑战. 前者为实现兵棋推演人机对抗的成功提供了技术基础, 后者则对当下人机对抗智能体决策技术提出了新的挑战.

1.1 兵棋推演问题

早期的兵棋推演一般指手工兵棋, 具有 200 多年的研究历史, 而随着信息技术与计算机性能的不断发展, 计算机兵棋因其简便、快速、逼真等特点, 成为目前兵棋推演的主流方向<sup>[18]</sup>. 2009 年, 彭春光等<sup>[16]</sup>对兵棋推演技术进行了综述, 指出兵棋主要研究人员决策与兵棋事件之间的因果关系. 2012 年,

王桂起等<sup>[15]</sup>概述了兵棋的概念、发展、分类以及应用, 并分析了兵棋的各组成要素以及国内外兵棋的研究现状. 2017 年, 胡晓峰等<sup>[6]</sup>对兵棋推演进行了全面的综述, 描述了兵棋推演的基本要素, 重点阐述了兵棋推演的关键在于模拟人的智能行为, 面临的难点为“假变真”、“粗变细”、“死变活”、“静变动”、“无变有”, 归结起来为“对战场态势的判断理解”和“对未来行动的正确决策处置”. 在此基础上, 作者展望了 AlphaGo 等技术对兵棋推演带来的新机遇.

不同于上述工作, 本文以人机对抗智能决策的角度切入, 分析兵棋推演中智能体研究的通用性智能决策技术与挑战.

1.2 策略游戏普遍挑战问题

回顾当前已获得一定人机对抗突破的典型决策环境(如雅达利、围棋、德州扑克和星际争霸), 可以得出一些基本结论: 人机对抗研究的重心已经从早期的单智能体决策环境(如雅达利)过渡到了多智能体决策环境(如围棋和星际争霸), 从回合制决策环境(如围棋)逐渐过渡到更贴近现实应用的复杂即时战略类决策环境(如星际争霸), 从完美信息博弈(如围棋)逐渐过渡到不完美信息博弈(如德州扑克和星际争霸), 从以树为基础的博弈算法(如围棋与德州扑克)过渡到以深度强化学习为基础的大规模机器学习算法. 典型的兵棋推演仿真环境一般由算子、地图、想定以及规则要素组成, 展现了红、蓝双方之间的博弈对抗. 与代表性策略游戏(如雅达利、围棋、德州扑克和星际争霸等)类似, 兵棋推演的智能体研究表现出策略游戏中智能体研究的普遍挑战性问题. 针对上述转变与各自博弈对抗环境的特点, 可以凝练抽取一些影响智能体设计与训练的关键因素, 如表 1 所述, 其中“√”表示存在, “×”表示不存在.

表 1 对决策带来挑战的代表性因素  
Table 1 Representative factors of challenge decision-making

| 游戏      | 雅达利 | 围棋 | 德州扑克 | 星际争霸 | 兵棋推演 |
|---------|-----|----|------|------|------|
| 不完美信息博弈 | √   | ×  | √    | √    | √    |
| 长时决策    | √   | √  | ×    | √    | √    |
| 策略非传递性  | ×   | √  | √    | √    | √    |
| 智能体协作   | ×   | ×  | ×    | √    | √    |
| 非对称环境   | ×   | ×  | ×    | ×    | √    |
| 随机性与高风险 | ×   | ×  | ×    | ×    | √    |

1) 不完美信息博弈. 不完美信息博弈是指没有参与者能够获得其他参与者的行动信息<sup>[19]</sup>, 即参与

者做决策时不知道或不完全知道自己所处的决策位置。相比于完美信息博弈, 不完美信息博弈挑战更大, 因为对于给定决策点, 最优策略的制定不仅仅与当下所处的子博弈相关。与德州扑克、星际争霸相似, 兵棋推演同样是不完美信息博弈, 红方或蓝方受限于算子视野范围、通视规则、掩蔽规则等, 需要推断对手的决策进而制定自己的策略。

2) 长时决策。相比于决策者仅做一次决策的单阶段决策游戏, 上述游戏属于序贯决策游戏<sup>[20]</sup>。以围棋为例, 决策者平均决策次数在 150 次, 相比于围棋, 星际争霸与兵棋推演的决策次数以千为单位。长时决策往往导致决策点数量指数级的增加, 使得策略空间复杂度变大, 过高的策略空间复杂度将带来探索与利用等一系列难题, 这对决策制定带来了极大的挑战。

3) 策略非传递性。对于任何策略  $v_t$  可战胜  $v_{t-1}$ ,  $v_{t+1}$  可战胜  $v_t$ , 有  $v_{t+1}$  可战胜  $v_{t-1}$ , 则认为策略之间存在传递性。一般情况下, 尽管部分决策环境存在必胜策略, 但在整个策略空间下都或多或少存在非传递性的部分, 即大多数博弈的策略不具备传递性<sup>[21]</sup>。例如在星际争霸和兵棋推演环境中, 策略难以枚举且存在一定的相互克制关系。策略非传递性会导致标准自博弈等技术手段难以实现智能体能力的迭代提升, 而当前经典的博弈算法 (如 Double oracle 等) 又往往难以处理大规模的博弈问题, 使得逼近纳什均衡策略极其困难。

4) 智能体协作。在多智能体合作环境中, 智能体间的协作将提升单个智能体的能力, 增加系统的鲁棒性, 适用于现实复杂的应用场景<sup>[22-24]</sup>。围棋与二人德州扑克参与方属于纯竞争博弈环境, 因此不存在多个智能体之间的协作。星际争霸与兵棋虽然也属于竞争博弈环境, 但是需要多兵力/算子之间配合获得多样化且高水平策略。将上述问题看作是单个智能体进行建模对求解是困难的, 可以建模为组队零和博弈, 队伍之间智能体相互协作, 最大化集体收益。针对组队零和博弈问题, 相比于二人零和博弈问题, 理论上相对匮乏。

除了上述代表性的、公认的影响因素之外, 其他一些对决策算法带来挑战性的因素有专家数据稀缺、奖励/反馈信号稀疏等。其中, 专家数据稀缺导致难以挖掘高水平人类数据, 也就难以构建知识库或进行高水平策略的分析, 尽管当前部分技术可以脱离人类数据进行高质量策略学习, 但是其面临着计算成本巨大、学习难度加剧等一系列问题。奖励/反馈信号稀疏, 意味着在策略的学习过程中无法获得即时有效的反馈, 也就导致策略的评估难度增大, 尤其是策略空间巨大时, 进行策略的搜索与评估将

变得尤为困难。

为应对上述挑战, 研究人员进行了大量的技术创新。例如, 在蒙特卡洛树搜索基础上引入深度神经网络实现博弈树剪枝、通过自博弈实现强化学习的围棋 AI AlphaGo 系列<sup>[2]</sup>, 在虚拟遗憾最小化算法基础上引入安全嵌套子博弈求解, 以及问题约简等技术的二人无限注德州扑克 AI Libratus<sup>[3]</sup>, 采用改进自博弈以及分布式强化学习的星际争霸 AI AlphaStar<sup>[4]</sup>。上述技术为相应决策问题的挑战性因素提出了可行的解决方案, 尽管兵棋推演存在上述挑战, 但相关技术基础已经具备, 可以指导兵棋推演的研究方向。

### 1.3 兵棋推演独特挑战问题

#### 1.3.1 非对称环境决策

传统的非对称信息是指某些行为人拥有但另一些行为人不拥有的信息, 本文的非对称从学习的角度考虑, 指的是游戏双方的能力水平或游戏平衡性。以围棋、星际争霸以及绝大多数游戏环境为例, 游戏设计者为保证游戏的体验以及促进人类选手竞技水平的提升, 往往保证游戏不同方具有相对均衡的能力。例如, 星际争霸游戏中包含了三个种族, 即人族、虫族和神族, 尽管不同种族具有截然不同的科技树、兵力类型等, 但是三个种族在能力上处于大致均衡的状态。

相比于星际争霸等, 兵棋推演中游戏是不平衡的。这不仅体现在红方和蓝方在兵力配备上的不同, 也体现在不同任务/想定下红方和蓝方的能力特长不同。以部分夺控战为例, 红方兵力水平一般弱于蓝方, 同时红方往往具有更好的侦察能力 (如红方配备巡飞弹算子), 而蓝方往往具有更强的进攻能力 (如配备更多的坦克算子)。这种严重的非对称性环境, 对于目前的学习算法提出了极大的挑战。

当前主流或改进的自博弈技术在智能体迭代过程中, 往往对每个参与智能体以对称的方式进行训练, 进而保证智能体能力在相互对抗的迭代过程中持续增长。但是在兵棋推演中, 红方和蓝方严重的非对称性环境, 使得直接采用相似设计难以保证弱勢方的训练, 需要设计更合理的迭代方式 (如启发式迭代) 保证相对较弱勢方的训练。另一方面, 在二人零和博弈中, 虽然弱勢方的纳什均衡策略可取, 但是如何根据对手的情况调整自己的策略, 以最大可能剥削或发现对手的漏洞并加以利用, 可能是需要考虑的重点问题。

#### 1.3.2 随机性与高风险决策

随机性与高风险主要体现在游戏的裁决中, 泛



指交战规则中随机影响因素以及对交战结果产生的影响. 裁决是游戏的重要组成部分, 在决定游戏的胜负规则之外, 明确定义了参与方在对抗过程中的交战结果. 例如在围棋中, 黑子包围白子之后, 需要将白子从棋盘中拿下, 即吃子. 在星际争霸环境中, 两队兵力对抗中, 血量为零的兵力将直接消失. 一般来说, 在围棋等棋类游戏中, 裁决不受随机因素的干扰, 即不具有随机性. 而在星际争霸环境中, 尽管不同兵力攻击产生的伤害数值是固定的, 但仍然受到少量随机因素的影响, 例如具有一定概率触发某项技能 (如闪避).

相比于上述游戏, 兵棋推演在所有攻击裁决过程中均受到随机因素的影响, 即随机性较高, 这主要是因为兵棋裁决一般遵循着“攻击等级确定、攻击等级修正、原始战果查询、最终战果修正”的基本流程. 在原始战果查询与最终战果修正中, 将基于骰子产生的随机数值 (两个骰子 1 ~ 12 点) 分别进行修正. 上述修正的结果差距较大, 可能产生压制甚至消灭对方班组的战果, 也有可能不产生任何效果. 更重要的是, 相比于其他即时战略类游戏 (如星际争霸), 兵力一旦消失, 将不能重新生成, 因此会造成极高的风险, 对于专业级选手, 兵力的消失往往意味着游戏的失败.

兵棋推演的随机性与高风险决策对于智能体的训练提出了极高的挑战. 反映在数据上, 环境的状态转移不仅受到其他算子以及不可见信息的影响, 也受到裁决的影响, 即状态转移高度不确定. 另一方面, 决策的高风险使得算子所处状态的值估计等具有高方差特性, 难以引导智能体的训练, 尤其是在评估上难以消除该随机性的情况下训练更加困难.

除了非对称环境决策和高风险与随机性决策之外, 兵棋推演因模拟真实对抗, 因此往往存在内部的与外部的干扰或影响. 例如对抗过程中工事或桥梁等建筑的损坏、恶劣天气导致的算子能力下降、人员的士气、想定之外额外的增援等. 在这些突发影响的情况下, 一般会对原有的策略造成较大的影响, 尤其是策略训练或设计之初未考虑这些因素时, 上述影响将会加剧. 当前兵棋推演中的这些偶然因素对智能体决策的影响的研究较少, 属于智能体的迁移与泛化等问题, 将是未来兵棋推演智能决策技术的研究方向之一.

总之, 兵棋推演作为一种人机对抗策略的验证环境, 其具备目前主流对抗环境 (如星际争霸等) 的挑战性问题, 因此当前的智能体开发技术与算法为获得有效的兵棋推演智能体提供了基础. 同时, 由于兵棋推演相对独特的不对称信息决策、更接近于

真实环境的随机性与高风险决策特点, 对当前人机对抗智能决策技术提出了新的挑战. 因此, 为了获得成功的兵棋推演智能体, 一方面需要对当前智能决策技术进行迁移, 解决面临的通用挑战在具体问题场景下的设计差异等问题. 例如策略非传递性这一挑战, 其在星际争霸与兵棋推演上表现不同, 则对自我博弈、群体博弈便提出了新的要求. 另一方面, 针对兵棋推演独特的挑战, 则需要结合上述设计进行必要的创新, 例如克服非对称环境决策下以自我博弈为技术路线的智能体迭代技术难以持续迭代的问题. 综上所述, 为了获得兵棋推演人机对抗突破, 需要结合前沿的智能决策技术, 并以此为基础, 克服兵棋推演的独特挑战性问题, 进而获得高水平的兵棋推演智能体.

## 2 兵棋智能决策技术研究现状

为应对兵棋推演的挑战性问题, 学者们提出了多种智能体研发与评测方法. 与围棋、星际争霸等主流游戏人机对抗智能体研发脉络类似 (如星际争霸从早期以知识规则为主, 中期以数据学习为主, 后期以联合知识与强化学习完成突破), 兵棋推演也经历了以知识驱动为主、以数据驱动为主以及以知识与数据混合驱动的研发历程. 兵棋的评测技术包含了智能体的定量与定性分析方法. 本节将重点阐述兵棋智能体研发的技术与框架, 同时对智能体的评估评测进行简述.

### 2.1 兵棋智能体研发技术与框架

当前, 智能体的研发技术与框架主要包含知识驱动、数据驱动和知识与数据混合驱动的兵棋推演智能体三类. 本节将分别阐述各个技术框架的研究进展.

#### 2.1.1 知识驱动的兵棋推演智能体

知识驱动的兵棋推演智能体研发利用人类推演经验形成知识库, 进而实现给定状态下的智能体决策<sup>[25]</sup>. 代表性的知识驱动框架为包以德循环 (Observation orientation decision action, OODA<sup>[26]</sup>), 其基本观点是通过观察、判断、决策和执行的循环过程实现决策, 如图 1 所示. 具体步骤为: 观察 (包括观察自己、环境以及对手) 实现信息的收集; 判断对应态势感知, 即对收集的数据进行分析、归纳以及总结获得当前的态与势; 决策对应策略的制定, 利用前面两步的结果实现最优策略的制定; 执行对应于具体的行动.

通过引入高水平人类选手的经验形成知识库, 可以一定程度规避前面所述的挑战性问题, 实现态

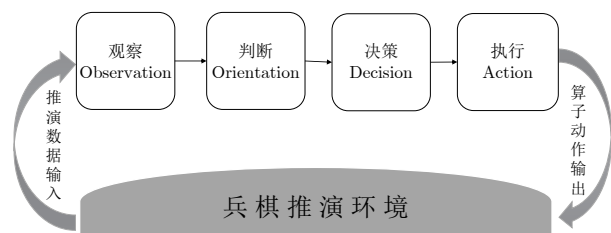


图1 包以德循环

Fig.1 OODA loop

势到决策的规则制定与编码。自 2017 年国内各类兵棋大赛举办以来, 每年都有数十甚至上百个参赛队伍进行机机对抗, 角逐的精英智能体将参与人机对抗以及人机混合对抗。为适应不同的想定以及进行人机协同, 目前绝大多数智能体为知识驱动型, 即依据人类选手的经验进行战法总结, 以行为树<sup>[27]</sup>、自动机<sup>[28]</sup>等框架实现智能体决策执行逻辑的编程。总的来说, 知识驱动型智能体研发依赖于人类推演经验与规律的总结, 实现相对简单, 不需要借助于大量的数据进行策略的训练与学习。

近年来, 通过模仿高水平选手的决策, 涌现出了一系列高水平知识驱动型智能体并开放对抗<sup>1</sup>。例如, 信息工程大学的“兵棋分队级 AI-微风 1.0”, 该智能体基于动态行为树框架, 在不同想定下实现了不同的战法战术库; 中国科学院自动化研究所的“兵棋群队级 AI-紫冬智剑 2.0”, 该智能体以 OODA 环为基本架构, 以敌情、我情以及地形等通用态势认识抽象状态空间, 以多层级任务行为认知抽象决策空间, 可以快速适应不同的任务/想定。目前, 部分智能体可以支撑人机混合对抗, 甚至在特定想定下达到了专业级选手水平。

### 2.1.2 数据驱动的兵棋推演智能体

随着 AlphaGo、AlphaStar 等智能体取得巨大成功, 以深度强化学习为基础进行策略自主迭代(如自博弈中每一轮的策略学习)成为当前的主流决策技术<sup>[29]</sup>并被成功应用于兵棋推演<sup>[30-31]</sup>, 其基本框架如图 2 所示。智能体以自博弈或改进的自博弈方式进行每一代智能体的迭代, 而每一代智能体采用强化学习的方式进行训练。对于强化学习来说, 智能体与环境进行交互收集状态、动作与奖赏等序列数据进行训练, 直至学习得到可以适应特定任务的策略。由于兵棋推演环境没有显式定义状态、动作与奖赏等的具体表现形式, 因此在应用强化学习的过程中, 首要的任务是进行上述基本要素的封装, 在此基础上便可以进行基本强化学习训练。

<sup>1</sup> <http://turingai.ia.ac.cn>

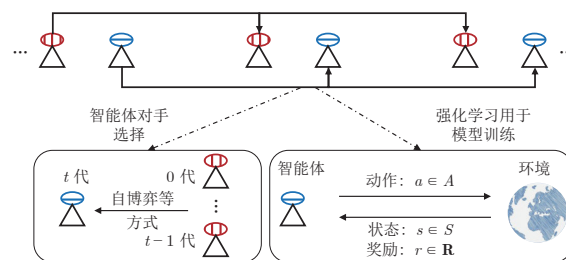


图2 自博弈加强化学习训练

Fig.2 Self-gaming and reinforcement learning

深度强化学习通过改进神经网络的设计, 可以一定程度缓解非完美信息与长时决策带来的挑战。例如通过增加认知网络结构(如记忆单元)<sup>[32-33]</sup>可以有效使用历史信息, 一定程度解决部分可观测状态下的决策问题。通过增加内在奖励驱动的环境建模网络<sup>[34]</sup>, 可以缓解长时决策尤其是奖励稀疏情况下强化学习的训练。自博弈尤其是改进的自博弈框架, 如星际争霸提出的带有优先级的虚拟自我对局与联盟训练有效缓解了策略非传递性的挑战, 并通过初期的强化学习网络监督训练初始化实现了对策略非传递性的进一步缓解。针对智能体协作, 学者们提出了大量的多智能体协同算法, 并通过奖励共享、奖励分配等实现了不同智能体的有效训练。关于非对称性与随机性与高风险, 据本文作者了解, 尚未有相关文献解决兵棋推演的上述挑战。

近年来, 一些学者将其他数据学习方式与强化学习进行结合以缓解端到端强化学习的困难。例如, 李琛等<sup>[30]</sup>将行动者-评论家框架引入兵棋推演并与规则结合进行智能体开发, 在简化想定(对称的坦克加步战车对抗)上进行了验证。张振等<sup>[31]</sup>将近端策略优化技术应用于智能体开发, 并与监督学习结合在智能体预训练基础上进行优化, 在简化想定(对称的两个坦克对抗)验证了策略的快速收敛。中国科学院自动化研究所提出的 AlphaWar<sup>2</sup> 引入监督学习与自博弈技术手段实现联合策略的学习, 保证了智能体策略的多样性, 一定程度缓解了兵棋推演的策略非传递性问题。2020 年, AlphaWar 在与专业级选手对抗过程中, 通过了图灵测试, 展现了强化学习驱动型兵棋推演智能体的技术优势。

另一方面, 分布式强化学习作为一种能够有效利用大规模计算资源加速强化学习训练的手段, 目前已成为数据驱动智能体研发的关键技术, 学者们提出了一系列算法, 在保证数据高效利用的同时也保证了策略训练的稳定性。例如, 2016 年, Mnih 等<sup>[35]</sup>提出异步优势动作评价算法, 实现了策略梯度算法

<sup>2</sup> [http://turingai.ia.ac.cn/ranks/wargame\\_list](http://turingai.ia.ac.cn/ranks/wargame_list)

的有效分布式训练. 2018 年, Horgan 等<sup>[36]</sup> 提出 APE-X 分布式强化学习算法, 对生成数据进行有效加权, 提升分布式深度 Q 网络训练效果. 2018 年, Espeholt 等<sup>[37]</sup> 提出重要性加权行动者-学习器框架 (Importance weighted actor-learner architecture, IMPALA) 算法实现了离策略分布式强化学习, 在高效数据产生的同时, 可以通过 V-Trace 算法进行离策略修正, 该技术被成功用于夺旗对抗<sup>[38]</sup>. 2020 年, Espeholt 等<sup>[39]</sup> 引入中心化模型统一前向, 进一步提升了 IMPALA 的分布式训练能力, 并被应用于星际争霸 AlphaStar 的训练中. 考虑到 IMPALA 的高效以及方便部署, 以 IMPALA 为代表的分布式强化学习已经成为兵棋智能体训练的常用算法. IMPALA 结构如图 3 所示, 包含学习器以及执行器两个主要部分. 其中学习器接收执行器获得的轨迹数据进行神经网络 (策略) 的参数更新, 执行器从学习器获得神经网络参数并与环境进行交互生成供训练用的轨迹数据, 这些数据通过数据队列的方式在学习器与执行器之间进行交换. 一般来说, 需要大量并发的执行器以在单位时间内产生足够多的数据. 针对兵棋推演红蓝双方, 则需要同时进行红蓝智能体的数据产生与策略参数更新. IMPALA 的代码实现相对简单, 可以方便地通过 Tensorflow、Pytorch、伯克利大学近期提出的 Ray<sup>[40]</sup> 框架完成.

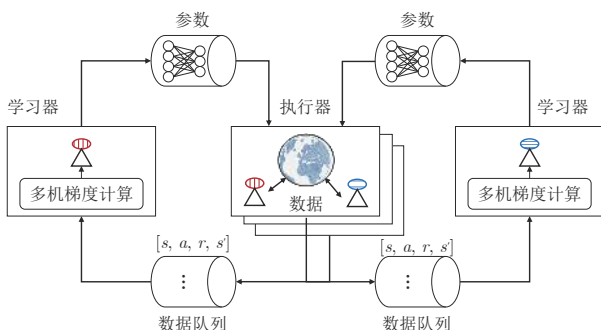


图 3 IMPALA 用于兵棋推演智能体训练  
Fig.3 IMPALA for training wargame agents

### 2.1.3 知识与数据混合驱动的智能体推演智能体

知识驱动智能体具有较强的可解释性, 但是受限于人类的推演水平. 与之相反, 基于数据驱动的智能体较少依赖人类推演经验, 可以通过自主学习的方式得到不同态势下的决策策略, 具有超越专业人类水平的潜力, 但是由于数据驱动的智能体推演智能体依赖数据以及深度神经网络, 其训练往往较为困难且决策算法缺乏可解释性.

为了有效融合知识驱动与数据驱动框架的优点, 避免各自的局限性, 目前越来越多学者试图将

两者进行结合<sup>[41]</sup>. 其中被关注较多的是将先验信息加入到学习过程中, 进而实现对机器学习模型的增强<sup>[42-44]</sup>. 在该类工作中, 知识或称为先验信息作为约束、损失函数等加入到学习的目标函数中实现一定程度的可解释性以及模型的增强. 近年来, Rueden 等<sup>[42]</sup> 进行了将知识融合到学习系统的综述并提出了知信机器学习概念, 从知识的来源、表示以及知识与机器学习管道的集成对现有方法进行了分类.

知识与数据混合驱动框架结合了两者的优势, 可以更好应对兵棋推演环境的挑战, 目前代表性的融合方式包括“加性融合”, 如图 4 所示, 即知识驱动与数据驱动各自做擅长的部分, 将其整合形成完整的智能体. 一般来说, 知识驱动善于处理兵棋推演前期排兵布阵问题, 因为该阶段往往缺乏环境的有效奖励设计. 另一方面, 紧急态势下的决策以及相对常识性的决策也可以由知识驱动完成, 以减少模型训练的探索空间. 数据驱动善于自动分析态势并作出决策, 更适用于进行兵棋推演中后期多样性策略的探索与学习. 此外, 一些难以用相对有限知识规则刻画的态势-决策也可由数据驱动完成. 黄凯奇等<sup>[45]</sup> 提出了一种融合知识与数据的人机对抗框架, 如图 5 所示, 该框架以 OODA 为基础, 刻画了决策不同阶段的关键问题, 不同问题可以通过数据驱动或知识驱动的方式进行求解.

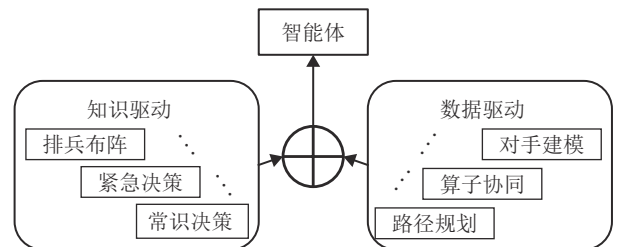


图 4 知识与数据驱动“加性融合”框架  
Fig.4 Additive fusion framework of knowledge and data driven

另一种代表性融合方式为“主从融合”, 如图 6 所示, 即以一方为主要框架, 另一方为辅助的融合方式. 在以知识驱动为主的框架中, 整体设计遵循知识驱动的方式, 在部分子问题或子模块上采用监督学习、进化学习等方式实现优化. 例如, 武警警官学院开发的分队/群队 AI “破晓星辰 2.0”<sup>3</sup> 在较为完善的人类策略库基础上, 结合蚁群或狼群等算法进行策略库优化, 以提升智能体的适应性. 在以数据驱动为主的框架下, 则采用数据驱动的改进自博弈强化学习的方式进行整体策略学习, 同时增加先验尤其是常识性约束. 例如, 将常识或人类经验

<sup>3</sup> <http://turingai.ia.ac.cn/notices/detail/116>



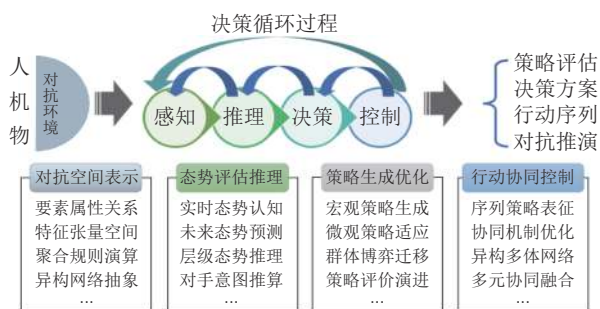
图5 人机对抗框架<sup>[45]</sup>Fig.5 Human-machine confrontation framework<sup>[45]</sup>

图6 知识与数据驱动“主从融合”框架

Fig.6 Principal and subordinate fusion framework of knowledge and data driven

作为神经网络选择动作的二次过滤以减少整体探索空间。

## 2.2 兵棋智能体评估评测及平台

智能体的评估涉及智能体整体能力与局部能力评估,同时开放的智能体评估平台将有效支撑智能体的能力测评与迭代。本节将从智能体评估算法与智能体评估开放平台进行介绍。

### 2.2.1 智能体评估算法

正确评估智能体策略的好坏,对于智能体的训练与能力迭代具有至关重要的作用。考虑到兵棋推演中策略的非传递性以及其巨大的策略空间问题,进行智能体准确评估具有巨大挑战。近年来,学者们提出了一系列评估算法,试图对智能体能力进行准确描述。经典的埃洛 (ELO) 算法<sup>[46]</sup>利用智能体之间的对抗结果,通过极大似然估计得到反映智能体能力的分值。例如围棋、星际争霸等对抗环境中的段位,就是基于 ELO 算法计算获得。Herbrich 等<sup>[47]</sup>提出 TrueSkill 算法,通过将对抗过程建立为因子关系图,借助于贝叶斯理论实现了多个智能体对抗中单一智能体能力的评估。考虑到 ELO 算法难以处理策略非传递性, Balduzzi 等<sup>[48]</sup>提出多维 ELO 算法,通过对非传递维度进行显式的近似,改善了胜率的预测问题。进一步, Omidshafiei 等<sup>[49]</sup>提出  $\alpha$ -rank 算法,基于 Markov-Conley 链,使用种群策略

进化的方法,对多种群中的策略进行排序,实现策略的有效评估。

定量评估之外,也可以通过专家评判的方式进行定性评估,实现对智能体单项能力的有效评估。例如,图7是庙算杯测试赛<sup>4</sup>中对智能体 AlphaWar 的评估,在定义的“武器运用”、“地形利用”、“兵力协同”、“策略高明”、“反应迅速”几个能力维度上采集选手和专家的评分,与测试赛排名第1位的人类选手进行了比较。

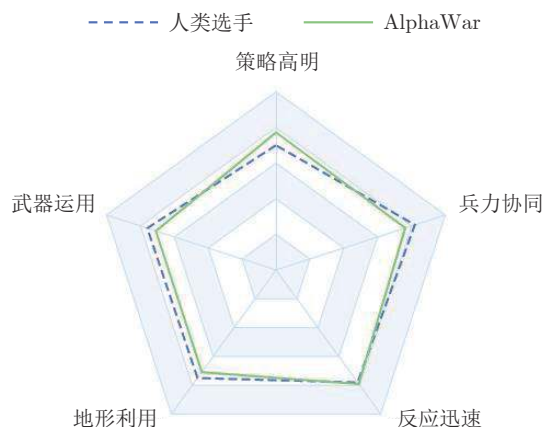


图7 智能体单项能力评估

Fig.7 Evaluation of specific capability of agents

### 2.2.2 智能体评估开放平台

为促进兵棋推演智能技术的发展,构建标准的评估评测平台至关重要,其可以实现广泛的兵棋智能体机机对抗、人机对抗甚至人机混合对抗<sup>[50]</sup>,这对兵棋推演评估评测平台提出了较高的要求,但也极大地促进了兵棋评估评测平台的建设与标准化。最近,中国科学院自动化研究所构建了人机对抗智能门户网站 <http://turingai.ia.ac.cn>,如图8所示。该平台以机器和人类对抗为途径,以博弈学习等为核心技术来实现机器智能快速学习进化的研究方向。平台提供兵棋推演智能体的机机对抗、人机对抗以及人机混合对抗测试,并支持智能体的多种评估评测。

## 3 兵棋智能决策技术的挑战

针对兵棋推演的智能技术研究现状,本节重点阐述不同技术框架存在的挑战性问题,引导学者们对相关问题进行深入研究。

### 3.1 知识驱动型技术挑战

知识驱动型作为智能体研发的主流技术之一,其依赖人类推演经验形成知识库,进而实现给定态

<sup>4</sup> <http://turingai.ia.ac.cn/bbs/detail/14/1/29>



图8 “图灵网”平台

Fig.8 Turing website platform

势下的智能体决策. 基于此, 知识驱动型智能体具有较强的可解释性, 但同样面临不可避免的局限, 即受限于人类本身的推演水平, 同时环境迁移与适应能力较差. 造成上述局限的根本原因在于缺乏高质量的知识库<sup>[51-52]</sup>, 实现知识建模、表示与学习<sup>[53]</sup>, 这也是目前知识驱动型技术的主要挑战. 知识库一般泛指专家系统设计所应用的规则集合, 其中规则所联系的事实及数据的全构成了知识库, 其具有层次化基本结构.

对于兵棋推演, 知识库最底层是“事实知识”(如算子机动力)等; 中间层是用来控制“事实”的知识(如规则等), 对应于兵棋中的微操等; 最顶层是“策略”, 用于控制中间层知识, 一般可以认为是规则的规则, 如图9所示. 兵棋推演中知识库构建过程最大的挑战便是顶层策略的建模, 面临着通用态势认知与推理困难的挑战. 胡晓峰等<sup>[6]</sup>指出兵棋推演需要突破作战态势智能认知瓶颈, 并提出战场态势层次不同, 对态势认知的要求和内容也不同. 尽管部分学者尝试从多尺度表达模型<sup>[54]</sup>、指挥决策智能体认知行为建模框架<sup>[55]</sup>以及基于OODA环框架下态势认知概念模型<sup>[56]</sup>等进行态势建模, 但是, 目前基于经典知识规划的智能体受限于对环境认识的正确性和完备程度, 表现相对呆板缺乏灵活应对能力, 不能很好地进行不确定环境边界下的意图估计与威胁评估等态势理解.

### 3.2 数据驱动型技术挑战

数据驱动型技术以深度强化学习为基础进行策略自主迭代, 从该角度出发解决兵棋推演智能体研发, 训练得到的智能体具有潜在的环境动态变化适应能力, 甚至有可能超越专业人类选手的水平, 涌现出新战法. 同样地, 为实现有效的智能体策略学习, 目前数据驱动型技术面临以下技术挑战: 自博弈与改进自博弈设计、多智能体有效协作、强化学

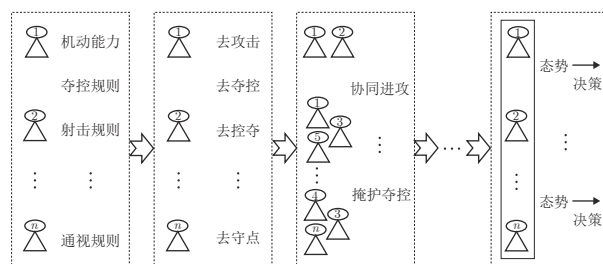


图9 兵棋推演知识库构建示例

Fig.9 Example of knowledge base construction for wargame

习样本效率较低. 其中, 自博弈与改进自博弈设计可以实现智能体能力的有效迭代提升, 多智能体有效协作将解决兵棋推演中的算子间协同(异步协同)问题, 而解决强化学习样本效率较低问题可以实现在可控计算资源与时间下的智能体训练.

1) 自博弈与改进自博弈设计. 在兵棋推演这个二人零和博弈问题下, 传统的博弈算法(如虚拟自我对局<sup>[57]</sup>、Double Oracle<sup>[58]</sup>等)难以适用于兵棋推演本身巨大的问题复杂度, 采用目前较为主流的自博弈或改进自博弈方式实现智能体能力的迭代, 成为一种可行的方案. 例如围棋游戏的AlphaGo系列<sup>[2]</sup>采用结合蒙特卡洛树搜索与深度神经网络的自博弈强化学习实现智能体能力的迭代. 星际争霸游戏的AlphaStar<sup>[4]</sup>则改进传统的虚拟自我对局, 提出带有优先级的虚拟自我对局并结合联盟训练进行智能体迭代. 具体地, AlphaStar引入主智能体、主要利用智能体以及联盟利用智能体, 并对不同的智能体采用不同的自博弈进行以强化学习为基础的参数更新. 总之, 尽管上述自博弈与一系列改进自博弈方法可以实现智能体的迭代, 但当前的设计多是启发式迭代方式, 兵棋推演的非对称环境等独特挑战是否适用, 有待验证与开展深入研究.

2) 多智能体有效协作. 协作环境下单个智能体的训练受到环境非平稳性的影响而变得不稳定<sup>[59-62]</sup>, 学者们提出了大量的学习范式以缓解该问题, 但仍然面临着智能体信用分配这一核心挑战, 即团队智能体在和环境交互时产生的奖励如何按照各个智能体的贡献进行合理分配以促进协作<sup>[63-65]</sup>. 目前, 一类典型的算法为Q值分解类算法, 即在联合Q值学习过程中按照单调性等基本假设将联合Q值分解为智能体Q值的联合, 进而实现信用隐式分配<sup>[66-68]</sup>. 例如, Sunehag等<sup>[66]</sup>率先提出此类算法将联合Q值分解为各个智能体Q值的加和. 在此基础上, Rashid等<sup>[67]</sup>基于单调性假设提出了更为复杂的Q值联合算法QMIX. 另外一类典型的信用分配算法借助于



差异奖励来实现显式奖励分配. 例如 Foerster 等<sup>[69]</sup>通过引入反事实的方法提出 COMA 以评估智能体的动作对联合智能体动作的贡献程度. 通过将夏普利值引入 Q 学习过程中, Nguyen 等<sup>[70]</sup>提出了 Shapley-Q 方法以实现“公平”的信用分配. 在兵棋推演环境中, 不同智能体原子动作执行耗时是不一样的, 导致智能体协作时的动作异步性, 如图 10 所示. 这种异步性使得智能体间的信用分配算法要求的动作同步性假设难以满足, 如何实现动作异步性下多智能体的有效协作仍然是相对开放的问题.

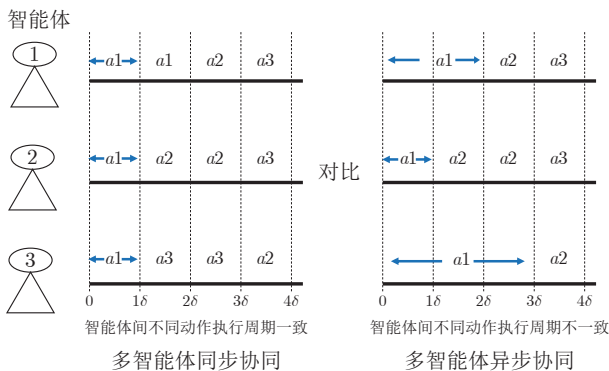


图 10 兵棋推演中的异步协同与同步协同对比

Fig. 10 Comparison between asynchronous cooperation and synchronous cooperation in wargame

3) 强化学习样本效率较低. 强化学习通过与环境交互试错的方式进行模型训练, 一般样本效率较低, 因此在复杂环境下智能体训练需要巨大计算资源. 例如, AlphaZero<sup>[71]</sup>采用了 5 000 一代 TPU 与 16 二代 TPU 进行智能体学习, AlphaStar<sup>[4]</sup>采用 192 TPU (8 核)、12 TPU (128 核) 与 50 400 CPU 实现群体博弈. 探索作为一种有效缓解样本效率低的手段<sup>[72]</sup>, 近年来受到了学者们的广泛关注, 并潜在适用具有巨大状态空间、稀疏奖励的兵棋推演环境中. 在单智能体强化学习中, 目前涌现了大量的探索类算法<sup>[72-74]</sup>, 如随机网络蒸馏<sup>[34]</sup>、Go Explore<sup>[73]</sup>等. 但多智能体的环境探索问题研究相对较少, 代表性方法有 MAVEN<sup>[75]</sup>、Deep-Q-DPP<sup>[76]</sup>、ROMA<sup>[77]</sup>等. 其中 MAVEN 通过在 QMIX 的基础上引入隐变量来实现多个联合 Q 值的学习, 进而完成环境的有效探索. Deep-Q-DPP 将量子物理中建模反费米子的行列式点过程引入多智能体探索中, 通过增加智能体行为的多样性来实现探索. 另一方面, ROMA 通过考虑智能体的分工, 让相同角色的单元完成相似的任务, 进而利用动作空间划分来实现环境高效探索. 上述算法在星际争霸微操等验证环境中取得了有效验证, 但是兵棋推演环境拥有更加庞大的状

态空间, 如何实现智能体异步动作下的环境高效探索, 对当前技术提出了新的要求.

### 3.3 知识与数据混合驱动型技术挑战

知识与数据混合驱动型, 相比于知识型与数据型, 可以有效融合两者的优点, 既具备对环境的适应能力, 涌现出超越高水平人类玩家的策略, 同时又具备一定可解释性, 实现可信决策. 在融合过程中面临知识与数据驱动本身的技术挑战外, 另一核心技术挑战在于融合方式, 即如何实现两者的有机融合<sup>[78]</sup>. 第 2.1.3 节的代表性的“加性融合”和“主从融合”, 可以实现知识与数据的一定程度融合, 但是何种融合方式更优目前并无定论. 另一方面, 探索更优的兵棋推演知识与数据融合思路是值得深入探索与研究的开放问题.

1) 加性融合的挑战. 在加性融合中, 知识驱动与数据驱动负责智能体不同的模块, 两者相加构成完整的智能体. 首先需要解决的问题是整个决策过程的模块化或解耦合. 目前, 兵棋推演中典型的加性融合框架依据对抗的进程实现两者的融合. 具体地, 在兵棋推演前期, 红、蓝双方将算子机动至中心战场, 为中后期对抗、夺控进而获胜创造条件. 该过程反馈较少, 且算子可探索空间巨大, 采用知识驱动的方式可以结合人类的领域知识, 相比于数据驱动更快速获得有效的策略. 在兵棋推演中期和后期, 双方发生剧烈、持续对抗, 策略的细小变化将产生不同的对抗结果, 该过程反馈较多, 且难以通过枚举的方式进行罗列, 采用数据驱动的方式可以结合已有的、自主生成的数据获得更加有效的策略, 但上述做法如何解耦合或定义两者的边界是困难的, 这不可避免引入专家的领域知识, 也将受限于专家们对问题认识的局限. 以 OODA 为基础的人机对抗框架<sup>[45]</sup>, 虽然给出了较为一般化的框架, 但是如何在兵棋推演中具体实现, 存在较大的不确定性. 另一方面, 知识驱动与数据驱动部分相互制约, 在设计或训练过程中势必受到彼此的影响. 例如数据驱动部分在迭代过程中受到知识驱动部分的限制, 这要求知识驱动或数据驱动部分在自我迭代的同时, 设计两者的交替迭代, 进而实现完整智能体能力的迭代提升. 上述设计与研究, 目前仍然是相对开放的问题.

2) 主从融合的挑战. 在主从融合中, 以知识驱动或数据驱动为主, 部分子问题以另一种方式手段进行解决. 在以数据驱动为主的框架中, 难点在于如何将知识或常识加入到深度学习或深度强化学习的训练中. 目前, 兵棋推演中典型的以数据驱动为主的框架, 与当前星际争霸、DOTA2 等智能体采

纳框架基本一致,即整体框架采用深度强化学习(自主博弈进行迭代)的数据驱动方式进行智能体学习,采用知识驱动的方式进行状态空间、动作空间尤其是奖励空间的设计.如何引入领域知识设计上述要素,将极大影响智能体的最终水平以及训练效率,因此需要对上述问题进行折中,保证智能体能力的同时尽可能引入更多的知识以提升训练效率.在以知识驱动为主的典型框架中,难点在于寻找适宜用学习进行解决的子问题,进而解决难以枚举或难以制定策略的场景.目前兵棋推演中,典型的以知识驱动为主的框架与早期星际争霸研究采纳的框架相似(如三星 SAIDA 智能体),即整体框架采用有限状态机、行为树的知识驱动方式进行智能体策略制定,采用数据驱动的方式进行如微操控制、敌方位置搜索的学习.例如采用经典的寻路算法<sup>[79]</sup>,实现临机路障等环境下的智能体机动设计,利用模糊系统方法实现兵棋进攻关键点<sup>[80]</sup>推理,基于关联分析模型进行兵棋推演武器效用的挖掘<sup>[81]</sup>.

### 3.4 评估评测技术挑战

当前,智能体的评估主要借助人机对抗的胜率进行智能体综合能力/段位的排名/估计.除此之外,兵棋推演一般建模为多智能体协作问题,因此单个智能体的能力评估将量化不同智能体的能力,在人机协作<sup>[82]</sup>中机的能力评估中占据重要的地位.另一方面,人机对抗中人对机的主观评价正逐渐成为一种智能体能力评估的重要补充.下面将分别介绍相关的挑战性问题.

1) 非传递性策略综合评估.多维 ELO 算法<sup>[48]</sup>在传统 ELO 的基础上,通过对非传递维度进行显式的近似,可以缓解非传递性策略胜率的预测问题,但是因为其依赖于 ELO 的计算方式,也就存在 ELO 本身对于对抗顺序依赖以及如何有效选取基准智能体等问题.对于兵棋推演这一面临严重策略非传递性的问题,目前的评估技术基于 ELO 或改进的 ELO,仍然具有较大的局限性.

2) 智能体协作中的单个智能体评估.基于经典的 ELO 算法,Jaeger 等<sup>[38]</sup>提出启发式的算法进行协作智能体中单个智能体的评估,但是该算法依赖于智能体能力的可加和假设,因此难以应用于兵棋推演环境,即算子之间的能力并非线性可加和.另一方面, TrueSkill 算法通过引入贝叶斯理论,实现了群体对抗中的某一选手的评估,但是其对时间不敏感,且往往会因为对抗选手的冗余出现评估偏差.因此,如何设计有效的评估算法实现协作智能体中的单个智能体的评估是当前的主要挑战之一.

3) 定性评估标准体系化.当前一些评估评测平台人为抽象了包括“武器使用”、“地形利用”等概念,构建评价维度,以实现人机对抗中人对智能体打分评测.上述概念主要启发于对指挥决策者能力的刻画,是面向应用刻画和评估智能体能力的重要维度<sup>[83-84]</sup>.但是,如何将智能体的评估体系与指挥控制能力维度进行对齐仍然是挑战性问题,需要指挥控制领域的研究人员与博弈决策领域的研究人员共同协作.

## 4 兵棋智能决策技术展望

近年来,人工智能技术在人机对抗领域获得了突飞猛进的发展,涌现出了一系列针对复杂博弈的智能体学习技术与算法,如分布式深度强化学习、自我博弈等.针对兵棋推演这一问题,早期研究者们致力于构建知识库或解决兵棋推演中的一些搜索或学习问题,如路径搜索<sup>[79]</sup>、意图估计<sup>[14]</sup>、关键点推理<sup>[80]</sup>、武器效用挖掘<sup>[81]</sup>等.最近,以深度学习尤其是深度强化学习为主的算法在兵棋推演中获得了长足的发展<sup>[11, 30-31]</sup>.但是,目前算法对于实现兵棋推演人机对抗突破尚存在一定的距离,当前的主流框架,如知识驱动、数据驱动以及知识与数据混合驱动的智能体学习算法,仍面临着一系列的挑战,如第 3 节所述.为缓解兵棋推演智能决策技术存在的挑战性问题,一些学者另辟蹊径,引入了新的理论、抽象约简问题等,以应对兵棋推演的人机对抗.

### 4.1 兵棋推演与博弈理论

博弈理论是研究多个利己个体之间的策略性交互而发展的数学理论,作为个体之间决策的一般理论框架,有望为兵棋人机对抗挑战突破提供理论支撑<sup>[85-88]</sup>.一般来说,利用博弈理论解决兵棋推演挑战,需要为兵棋推演问题定义博弈解,并对该解进行计算<sup>[89]</sup>.对于兵棋推演中典型的二人零和博弈问题,可以采用纳什均衡解.但是,纳什均衡解作为一种相对保守的解,并非在所有场合都适用.考虑到兵棋推演的严重非对称性,纳什均衡解对于较弱势方可能并不合适.因此,如何改进纳什均衡解(例如以纳什均衡解为基础进行对对手剥削解的迁移)是需要研究的关键问题.

在博弈求解问题上,早期相对成熟的求解方法包括线性规划、虚拟自我对局<sup>[57]</sup>、策略空间回应 Oracle<sup>[90]</sup>、Double oracle<sup>[58]</sup>、反事实遗憾最小化<sup>[91]</sup>等.但是,上述纳什均衡解(或近似纳什均衡解)优化方法一般只能处理远低于兵棋推演复杂度的博弈环境,而目前主流的用于星际争霸等问题的基于启发式设计的改进自博弈迭代,尚缺乏对纳什均衡逼近



近的理论保证. 因此, 针对兵棋推演这一具有高复杂度的不完美信息博弈问题, 如何将深度强化学习技术有效地纳入可逼近纳什均衡解的计算框架, 或提出更有效/易迭代的均衡逼近框架, 来实现兵棋推演解的计算, 仍然是开放性问题.

总之, 尽管博弈理论为兵棋推演的人机对抗挑战提供了理论指导, 但如何借助于该理论实现兵棋推演人机对抗的突破, 仍然是相对开放性的问题, 需要学者们更深入地进行研究.

#### 4.2 兵棋推演与大模型

近年来, 大模型(预训练模型)在自然语言处理领域获得了飞速发展<sup>[92-93]</sup>. 例如, 2020年, OpenAI发布的 GPT-3 模型参数规模达到 1 750 亿个<sup>[94]</sup>, 可以作为零样本或小样本学习模型提升自然语言处理任务的性能, 如文本分类、文本生成与对话生成. 2021年, 中国科学院自动化研究所在世界人工智能大会上, 发布了视觉、文本、语音三模态大模型, 其具备跨模态理解与生成能力<sup>5</sup>. 大模型作为通用人工智能的一种探索路径, 一般需要大量的数据进行模型预训练, 发展至今, 其具有重要的学术研究价值与广泛的应用前景.

兵棋推演提供多种任务/想定, 理论上可以有大量不同的训练环境, 深度强化学习与环境交互试错的学习机制使得大模型训练的数据问题得以缓解. 但是, 如何针对兵棋推演训练大模型, 使得其在不同的兵棋对抗任务中可以快速适应, 仍然面临各种挑战, 如图 11 所示. 首先, 兵棋推演没有较为通用的训练目标或优化目标(如自然语言处理任务), 尤其是不同规模的对抗任务差异较大, 因此如何设计该大模型的优化目标是需要解决的首要问题, 这涉及强化学习中动作空间、奖励空间等多项要素的深入考虑.

另一方面, 兵棋推演包含异质且异步协同的智能体, 不同任务下需要协同的智能体在数量、类型上有所差距, 这就要求大模型在训练过程中既能解耦不同智能体之间的训练, 同时可以建立有效的协同机制, 实现智能体之间的协同. 尽管可以采用智能体共享奖励、神经网络独立训练的框架, 但该设计过于简单, 难以有效实现智能体协同时的信用分配等挑战性问题. 总之, 如何设计大模型下多智能体训练以适应具有较大差异的兵棋推演任务, 是需要重点研究的问题之一.

最后, 在自博弈过程中进行大模型的训练, 需要适应不同规模(兵棋推演天然存在连队级、群队

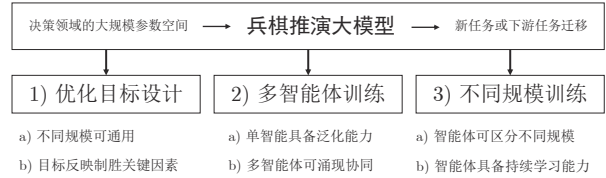


图 11 兵棋推演大模型训练挑战

Fig. 11 Challenges of training big model for wargame

级、旅队级等规模) 以及同规模下不同任务难度的对抗, 这对大模型的训练提出了新的挑战. 自步学习<sup>[95]</sup>的范式提供了智能体由易到难的逐步训练框架, 但如何定义兵棋推演不同任务难度是启发式的. 另一方面, 要求智能体在更难任务训练时不能遗忘对已训练任务的记忆, 这也需要持续学习<sup>[96]</sup>等前沿技术手段的引入.

#### 4.3 兵棋推演关键问题抽象

星际争霸完整游戏的人机对抗挑战突破之前, 学者们设计了包括敌方意图识别<sup>[97]</sup>、微操控制(多智能体协同)<sup>[98-100]</sup>等在内的关键子任务以促进智能决策技术的发展. 针对兵棋推演问题, 为引领技术突破进而反馈解决兵棋人机对抗挑战, 迫切需要对兵棋推演中的关键问题进行抽象、约简, 在保证约简的问题能够表征原始问题的重要特征前提下, 在约简的问题中进行求解.

基于上述考虑, 本文提出排兵布阵与算子异步协同对抗两个约简问题. 需要指出的是, 在问题约简过程中, 不可避免对兵棋推演环境等要素的规则进行简化, 甚至脱离兵棋推演本身的任务目的或者导向属性, 但相关问题的约简与抽象一定程度反映了兵棋推演智能体决策的核心挑战, 将极大促进学者们对相关问题的研究.

1) 排兵布阵反映了决策者在未知对手如何决策的前提下, 采取何种规划或兵力选择, 可以最大化自己的收益, 代表性环境有炉石传说卡牌类游戏, 即如何布置自己的卡牌以在后期积累优势获得最大化利益. 其挑战在于未知对手如何规划条件下, 实现己方规划, 该问题因为缺乏验证环境, 目前研究较少.

兵棋推演前期, 红方或蓝方基于未知的对手信息布局自己的兵力, 该布局一定程度决定了后期的对抗成败. 该过程因为缺少环境的显式反馈, 无法度量何种排兵布阵能够最大限度利用地形、能够最大化攻击等, 也就难以评估何种兵力布置最优. 基于上述原因, 本文设计了排兵布阵简化问题, 如图 12 所示, 即在一个简化地图中, 红方与蓝方各占有一部分区域进行兵力放置, 同时红方与蓝方之间具有

<sup>5</sup> [http://www.cas.cn/syky/202107/t20210712\\_4798152.shtml](http://www.cas.cn/syky/202107/t20210712_4798152.shtml)



一定距离间隔,考虑红方与蓝方不能移动且兵力放置之后自动进行裁决.

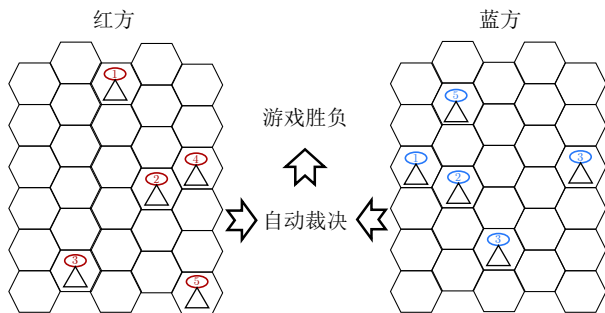


图 12 排兵布阵问题示例

Fig. 12 Example for problem of arranging arms

需要指出的是,上述简化环境对兵棋推演本身做了极大的简化,更多是从算法研究的角度出发.在研究兵力放置过程中,可以由简单到复杂进行调整,以契合兵棋推演问题本身,包括兵棋推演的目的加入(如夺控)、地形设置(如高程)等.

2) 算子协同对抗是多智能体相关问题的重要组成部分,目前相关领域已经开放了大量的智能体协同对抗环境,如星际争霸微操、捉迷藏等<sup>[22-24]</sup>.值得注意的是,目前在大多数环境中,不同算子之间协同是同步的,即智能体的动作执行周期一致.以此为基础,学者们提出了大量的算法实现有效的算子间协同<sup>[75, 100-101]</sup>.但是当不同智能体的动作执行周期不一致时,便导致异步协同问题,如兵棋推演的对抗便属于异步协同对抗,因为相关环境的缺乏,当前的研究相对较少.

兵棋推演中期和后期,红方与蓝方进行对抗,为评估智能体的接敌能力实现算子之间异步动作的有效协同,本文设计算子异步协同对抗简化问题.如图 13 所示.在一个简化的相对较小的地图上,不考虑复杂地形、复杂交战规则以及兵棋推演任务约束等因素,红方与蓝方在各自的起始位置出发进行对抗,算子可选动作包括机动(6 个方向与停止)与射击(对方算子).由于不同算子机动能力的差异,重点为领域提供多智能体异步协作的评估环境.

同排兵布阵问题一样,简化更多从验证算法性能的角度入手.在研究算子异步协同对抗过程中,可以对任务的难度进行调整,如对地图进行调整,包括设置高程、增加特殊地形等.

为了促进上述问题的深入研究,在约简问题设计上需要保证以下三点:

1) 与 OpenAI Gym<sup>6</sup> 一致的领域认可的环境接口,供智能体与环境交互进行策略的学习;

<sup>6</sup> <http://gym.openai.com>

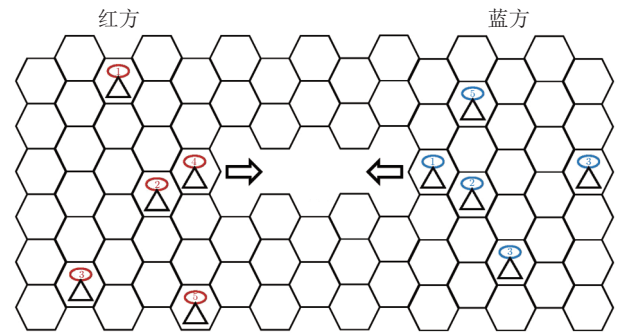


图 13 异步协同问题示例

Fig. 13 Example for problem of asynchronous multi-agent cooperation

2) 提供不同难度等级的内置智能体,供算法研究人员进行算法验证与算法间比较;

3) 完全开放的底层源码,进而支持自博弈等主流技术以及人机对抗.

## 5 结束语

星际争霸人机对抗挑战的成功,标志着智能决策技术在高复杂不完美信息博弈中的突破,之后迫切需要新的人机对抗环境以牵引智能决策技术的发展.兵棋推演因其非对称信息决策以及随机性与高风险决策等挑战性问题,潜在成为下一个人机对抗热点.本文详细分析了兵棋推演智能体的研究挑战,尤其是相比于其他博弈环境的独特挑战性问题.在此基础上,梳理了兵棋推演智能决策技术的研究现状,包括智能体研发技术框架以及智能体评估评测技术.然后,指出了当前技术的挑战,并展望兵棋推演智能决策技术的发展趋势.通过本文,将启发学者对兵棋推演关键问题的研究,进而产生实际应用价值.

## References

- 1 Campbell M, Hoane A J J, Hsu F H. Deep blue. *Artificial Intelligence*, 2002, **134**(1-2): 57-83
- 2 Silver D, Huang A, Maddison C J, Guez A, Sifre L, Van Den Driessche G, et al. Mastering the game of Go with deep neural networks and tree search. *Nature*, 2016, **529**(7587): 484-489
- 3 Brown N, Sandholm T. Superhuman AI for heads-up no-limit poker: Libratus beats top professionals. *Science*, 2018, **359**(6374): 418-424
- 4 Vinyals O, Babuschkin I, Czarnecki W M, Mathieu M, Dudzik A, Chung J, et al. Grandmaster level in StarCraft II using multi-agent reinforcement learning. *Nature*, 2019, **575**(7782): 350-354
- 5 Ye D H, Chen G B, Zhang W, Chen S, Yuan B, Liu B, et al. Towards playing full MOBA games with deep reinforcement learning. In: *Proceedings of the Advances in Neural Information Processing Systems 33. Virtual Event: MIT Press*, 2020.
- 6 Hu Xiao-Feng, He Xiao-Yuan, Tao Jiu-Yang. AlphaGo's breakthrough and challenges of wargaming. *Science & Technology Review*, 2017, **35**(21): 49-60 (胡晓峰, 贺筱媛, 陶九阳. AlphaGo 的突破与兵棋推演的挑战. *科技导报*, 2017, **35**(21): 49-60)

- 7 Hu Xiao-Feng, Qi Da-Wei. Intelligent wargaming system: Change needed by next generation need to be changed. *Journal of System Simulation*, 2021, **33**(9): 1997–2009  
(胡晓峰, 齐大伟. 智能化兵棋系统: 下一代需要改变的是什么. 系统仿真学报, 2021, **33**(9): 1997–2009)
- 8 Wu Lin, Hu Xiao-Feng, Tao Jiu-Yang, He Xiao-Yuan. Wargaming eco-system for intelligence growing. *Journal of System Simulation*, 2021, **33**(9): 2048–2058  
(吴琳, 胡晓峰, 陶九阳, 贺筱媛. 面向智能成长的兵棋推演生态系统. 系统仿真学报, 2021, **33**(9): 2048–2058)
- 9 Xu Jia-Le, Zhang Hai-Dong, Zhao Dong-Hai, Ni Wan-Cheng. Tactical maneuver strategy learning from land wargame replay based on convolutional neural network. *Journal of System Simulation*, 2022, **34**(10): 2181–2193  
(徐佳乐, 张海东, 赵东海, 倪晚成. 基于卷积神经网络的陆战兵棋战术机动策略学习. 系统仿真学报, 2022, **34**(10): 2181–2193)
- 10 Moy G, Shekh S. The application of AlphaZero to wargaming. In: Proceedings of the 32nd Australasian Joint Conference on Artificial Intelligence. Adelaide, Australia: 2019. 3–14
- 11 Wu K, Liu M, Cui P, Zhang Y. A training model of wargaming based on imitation learning and deep reinforcement learning. In: Proceedings of the Chinese Intelligent Systems Conference. Beijing, China: 2022. 786–795
- 12 Hu Gen-Sheng, Zhang Qian-Qian, Ma Chao-Zhong. Outlier data mining of the war game system. *Journal of Information Engineering University*, 2020, **21**(3): 373–377  
(胡良胜, 张倩倩, 马朝忠. 兵棋推演系统中的异常数据挖掘方法. 信息工程大学学报, 2020, **21**(3): 373–377)
- 13 Zhang Jin-Ming. Trajectory clustering algorithm of wargaming system based on space-time and combat grouping. *Journal of System Simulation*, 2016, **28**(8): 1748–1756  
(张锦明. 运用栅格矩阵建立兵棋地图的地形属性. 系统仿真学报, 2016, **28**(8): 1748–1756)
- 14 Chen L, Liang X, Feng Y, Zhang L, Yang J, Liu Z. Online intention recognition with incomplete information based on a weighted contrastive predictive coding model in wargame. *IEEE Transaction on Neural Networks and Learning Systems*, DOI: 10.1109/TNNLS.2022.3144171
- 15 Wang Gui-Qi, Liu Hui, Zhu Ning. A survey of war games technology. *Ordnance Industry Automation*, 2012, **31**(8): 38–41, 45  
(王桂起, 刘辉, 朱宁. 兵棋技术综述. 兵工自动化, 2012, **31**(8): 38–41, 45)
- 16 Peng Chun-Guang, Zhao Xin-Ye, Liu Bao-Hong, Huang Ke-Di. The technology of wargaming: An overview. In: Proceedings of the 14th Chinese Conference on System Simulation Technology & Application. Hefei, China: 2009. 366–370  
(彭春光, 赵鑫业, 刘宝宏, 黄柯棣. 兵棋推演技术综述. 第 14 届系统仿真技术及其应用学术会议. 合肥, 中国: 2009. 366–370)
- 17 Cao Zhan-Guang, Tao Shuai, Hu Xiao-Feng, He Lyu-Long. Abroad wargaming deduction and system research. *Journal of System Simulation*, 2021, **33**(9): 2059–2065  
(曹占广, 陶帅, 胡晓峰, 何吕龙. 国外兵棋推演及系统研究进展. 系统仿真学报, 2021, **33**(9): 2059–2065)
- 18 Si Guang-Ya, Wang Yan-Zheng. Challenges and reflection on next-generation large-scale computer wargame system. *Journal of System Simulation*, 2021, **33**(9): 2010–2016  
(司光亚, 王艳正. 新一代大型计算机兵棋系统面临的挑战与思考. 系统仿真学报, 2021, **33**(9): 2010–2016)
- 19 Ganzfried S, Sandholm T. Game theory-based opponent modeling in large imperfect-information games. In: Proceedings of the 10th International Conference on Autonomous Agents and Multi-agent Systems. Taipei, China: 2011. 533–540
- 20 Littman M L. Algorithms for Sequential Decision Making [Ph.D. dissertation], Brown University, USA, 1996
- 21 Nieves N P, Yang Y D, Slumbers O, Mguni D H, Wen Y, Wang J. Modelling behavioural diversity for learning in open-ended games. In: Proceedings of the 38th International Conference on Machine Learning. Vienna, Austria: 2021. 8514–8524
- 22 Huang X, Leng S, Maharjan S, Zhang Y. Multi-agent deep reinforcement learning for computation offloading and interference coordination in small cell networks. *IEEE Transactions on Vehicular Technology*, 2021, **70**(9): 9282–9293
- 23 Baker B, Kanitscheider I, Markov T M, Wu Y, Powell G, McGrew B, et al. Emergent tool use from multi-agent autocurricula. In: Proceedings of the 8th International Conference on Learning Representations. Addis Ababa, Ethiopia: 2020.
- 24 Liu I J, Jain U, Yeh R A, Schwing A G. Cooperative exploration for multi-agent deep reinforcement learning. In: Proceedings of the 38th International Conference on Machine Learning. Vienna, Austria: 2021. 6826–6836
- 25 Zhou Zhi-Jie, Cao You, Hu Chang-Hua, Tang Shuai-Wen, Zhang Chun-Chao, Wang Jie. The interpretability of rule-based modeling approach and its development. *Acta Automatica Sinica*, 2021, **47**(6): 1201–1216  
(周志杰, 曹友, 胡昌华, 唐帅文, 张春潮, 王杰. 基于规则的建模方法的可解释性及其发展. 自动化学报, 2021, **47**(6): 1201–1216)
- 26 Révay M, Liska M. OODA loop in command & control systems. In: Proceedings of the Communication and Information Technologies. Vysoke Tatry, Slovakia: 2017.
- 27 Nicolau M, Perez-Liebana D, O'Neill M, Brabazon A. Evolutionary behavior tree approaches for navigating platform games. *IEEE Transactions on Computational Intelligence and AI in Games*, 2017, **9**(3): 227–238
- 28 Najam-ul-Islam M, Zahra F T, Jafri A R, Shah R, Hassan M U, Rashid M. Auto-implementation of parallel hardware architecture for Aho-Corasick algorithm. *Design Automation for Embedded System*, 2022, **26**: 29–53
- 29 Cui Wen-Hua, Li Dong, Tang Yu-Bo, Liu Shao-Jun. Framework of wargaming decision-making methods based on deep reinforcement learning. *National Defense Technology*, 2020, **41**(2): 113–121  
(崔文华, 李东, 唐宇波, 柳少军. 基于深度强化学习的兵棋推演决策方法框架. 国防科技, 2020, **41**(2): 113–121)
- 30 Li Chen, Huang Yan-Yan, Zhang Yong-Liang, Chen Tian-De. Multi-agent decision-making method based on Actor-Critic framework and its application in wargame. *Systems Engineering and Electronics*, 2021, **43**(3): 755–762  
(李琛, 黄炎焱, 张永亮, 陈天德. Actor-Critic 框架下的多智能体决策方法及其在兵棋上的应用. 系统工程与电子技术, 2021, **43**(3): 755–762)
- 31 Zhang Zhen, Huang Yan-Yan, Zhang Yong-Liang, Chen Tian-De. Battle entity confrontation algorithm based on proximal policy optimization. *Journal of Nanjing University of Science and Technology*, 2021, **45**(1): 77–83  
(张振, 黄炎焱, 张永亮, 陈天德. 基于近端策略优化的作战实体博弈对抗算法. 南京理工大学学报, 2021, **45**(1): 77–83)
- 32 Qin Chao, Gao Xiao-Guang, Wan Kai-Fang. Deep spatio-temporal convolutional long-short memory network. *Acta Automatica Sinica*, 2020, **46**(3): 451–462  
(秦超, 高晓光, 万开方. 深度卷积记忆网络时空数据模型. 自动化学报, 2020, **46**(3): 451–462)
- 33 Chen Wei-Hong, An Ji-Yao, Li Ren-Fa, Li Wan-Li. Review on deep-learning-based cognitive computing. *Acta Automatica Sinica*, 2017, **43**(11): 1886–1897  
(陈伟宏, 安吉尧, 李仁发, 李万里. 深度学习认知计算综述. 自动化学报, 2017, **43**(11): 1886–1897)
- 34 Burda Y, Edwards H, Storkey A J, Klimov O. Exploration by random network distillation. In: Proceedings of the 7th International Conference on Learning Representations. New Orleans, USA: 2019.
- 35 Mnih V, Badia A P, Mirza M, Graves A, Harley T, Lillicrap T P, et al. Asynchronous methods for deep reinforcement learning. In: Proceedings of the 33rd International Conference on Machine Learning. New York, USA: 2016. 1928–1937
- 36 Horgan D, Quan J, Budden D, Barth-Maron G, Hessel M, Van Hasselt H, et al. Distributed prioritized experience replay. In:

- Proceedings of the 6th International Conference on Learning Representations. Vancouver, Canada: 2018.
- 37 Espeholt L, Soyer H, Munos R, Simonyan K, Mnih V, Ward T, et al. IMPALA: Scalable distributed deep-RL with importance weighted actor-learner architectures. In: Proceedings of the 35th International Conference on Machine Learning. Stockholm, Sweden: 2018. 1407–1416
  - 38 Jaderberg M, Czarnecki W M, Dunning I, Marris L, Lever G, Castañeda A G, et al. Human-level performance in 3D multi-player games with population-based reinforcement learning. *Science*, 2019, **364**(6443): 859–865
  - 39 Espeholt L, Marinier R, Stanczyk P, Wang K, Michalski M. SEED RL: Scalable and efficient deep-RL with accelerated central inference. In: Proceedings of the 8th International Conference on Learning Representations. Addis Ababa, Ethiopia: 2020.
  - 40 Moritz P, Nishihara R, Wang S, Tumanov A, Liaw R, Liang E, et al. Ray: A distributed framework for emerging AI applications. In: Proceedings of the 13th USENIX Conference on Operating Systems Design and Implementation. Carlsbad, USA: 2018. 561–577
  - 41 Pu Zhi-Qiang, Yi Jian-Qiang, Liu Zhen, Qiu Teng-Hai, Sun Jin-Lin, Li Fei-Mo. Knowledge-based and data-driven integrating methodologies for collective intelligence decision making: A survey. *Acta Automatica Sinica*, 2022, **48**(3): 627–643  
(蒲志强, 易建强, 刘振, 丘腾海, 孙金林, 李非墨. 知识和数据协同驱动的群体智能决策方法研究综述. 自动化学报, 2022, **48**(3): 627–643)
  - 42 Rueden L V, Mayer S, Beckh K, Georgiev B, Giesselbach S, Heese R, et al. Informed machine learning: A taxonomy and survey of integrating prior knowledge into learning systems. *IEEE Transactions on Knowledge and Data Engineering*, 2021, (5): 1–19
  - 43 Hartmann G, Shiller Z, Azaria A. Deep reinforcement learning for time optimal velocity control using prior knowledge. In: Proceedings of the 31st International Conference on Tools With Artificial Intelligence. Portland, USA: 2019. 186–193
  - 44 Zhang P, Hao J Y, Wang W X, Tang H Y, Ma Y, Duan Y H, et al. KoGuN: Accelerating deep reinforcement learning via integrating human suboptimal knowledge. In: Proceedings of the 29th International Joint Conference on Artificial Intelligence. Virtual Event: 2020. 2291–2297
  - 45 Huang Kai-Qi, Xing Jun-Liang, Zhang Jun-Ge, Ni Wan-Cheng, Xu Bo. Intelligent technologies of human-computer gaming. *SCIENTIA SINICA: Informations*, 2020, **50**(4): 540–550  
(黄凯奇, 兴军亮, 张俊格, 倪晚成, 徐博. 人机对抗智能技术. 中国科学: 信息科学, 2020, **50**(4): 540–550)
  - 46 Elo A E. *The Rating of Chess Players, Past and Present*. London: Batsford, 1978.
  - 47 Herbrich R, Minka T, Graepel T. TrueSkill (TM): A Bayesian skill rating system. In: Proceedings of the 19th International Conference on Neural Information Processing Systems. Vancouver, Canada: 2006. 569–576
  - 48 Balduzzi D, Tuyls K, Perolat J, Graepel T. Re-evaluating evaluation. In: Proceedings of the 32nd International Conference on Neural Information Processing Systems. Montréal, Canada: 2018. 3272–3283
  - 49 Omidshafiei S, Papadimitriou C, Piliouras G, Tuyls K, Rowland M, Lespiau J B, et al.  $\alpha$ -rank: Multi-agent evaluation by evolution. *Scientific Reports*, 2019, **9**(1): Article No. 9937
  - 50 Tang Yu-Bo, Shen Bi-Long, Shi Lei, Yi Xing. Research on the issues of next generation wargame system model engine. *Journal of System Simulation*, 2021, **33**(9): 2025–2036  
(唐宇波, 沈弼龙, 师磊, 易星. 下一代兵棋系统模型引擎设计问题研究. 系统仿真学报, 2021, **33**(9): 2025–2036)
  - 51 Ji S, Pan S, Cambria E, Marttinen P, Yu P. A survey on knowledge graphs: Representation, acquisition, and applications. *IEEE Transactions on Neural Networks and Learning Systems*, 2022, **33**(2): 494–514
  - 52 Wang Z, Zhang J W, Feng J L, Chen Z. Knowledge graph embedding by translating on hyperplanes. In: Proceedings of the 28th AAAI Conference on Artificial Intelligence. Québec, Canada: 2014. 1112–1119
  - 53 Wang Bao-Kui, Wu Lin, Hu Xiao-Feng, He Xiao-Yuan, Guo Sheng-Ming. Operations command behavior knowledge representation learning method based on sequential graph. *Systems Engineering and Electronics*, 2020, **42**(11): 2520–2528  
(王保魁, 吴琳, 胡晓峰, 贺筱媛, 郭圣明. 基于时序图的作战指挥行为知识表示学习方法. 系统工程与电子技术, 2020, **42**(11): 2520–2528)
  - 54 Liu Song, Wu Zhi-Qiang, You Xiong, Zhang Xin, Wang Xue-Feng. Multi-scale expression of integrated battlefield situation based on wargaming. *Journal of Geomatics Science and Technology*, 2012, **29**(5): 382–385, 390  
(刘嵩, 武志强, 游雄, 张欣, 王雪峰. 基于兵棋推演的综合战场态势多尺度表达. 测绘科学技术学报, 2012, **29**(5): 382–385, 390)
  - 55 He Xiao-Yuan, Guo Sheng-Ming, Wu Lin, Li Dong, Xu Xiao, Li Li. Modeling research of cognition behavior for intelligent wargaming. *Journal of System Simulation*, 2021, **33**(9): 2037–2047  
(贺筱媛, 郭圣明, 吴琳, 李东, 许霄, 李丽. 面向智能化兵棋的认知行为建模方法研究. 系统仿真学报, 2021, **33**(9): 2037–2047)
  - 56 Zhu Feng, Hu Xiao-Feng, Wu Lin, He Xiao-Yuan, Lv Xue-Zhi, Liao Ying. From situation cognition stepped into situation intelligent cognition. *Journal of System Simulation*, 2018, **30**(3): 761–771  
(朱丰, 胡晓峰, 吴琳, 贺筱媛, 吕学志, 廖鹰. 从态势认知走向态势智能认知. 系统仿真学报, 2018, **30**(3): 761–771)
  - 57 Heinrich J, Lanctot M, Silver D. Fictitious self-play in extensive-form games. In: Proceedings of the 32nd International Conference on Machine Learning. Lille, France: 2015. 805–813
  - 58 Adam L, Horcík R, Kasl T, Kroupa T. Double oracle algorithm for computing equilibria in continuous games. In: Proceedings of the 35th AAAI Conference on Artificial Intelligence. Virtual Event: 2021. 5070–5077
  - 59 Nguyen T T, Nguyen N D, Nahavandi S. Deep reinforcement learning for multi-agent systems: A review of challenges, solutions, and applications. *IEEE Transactions on Cybernetics*, 2020, **50**(9): 3826–3839
  - 60 Zhang K Q, Yang Z R, Başar T. Multi-agent reinforcement learning: A selective overview of theories and algorithms. *Handbook of Reinforcement Learning and Control*, 2021: 321–384
  - 61 Shi Wei, Feng Yang-He, Cheng Guang-Quan, Huang Hong-Lan, Huang Jin-Cai, Liu Zhong, et al. Research on multi-aircraft cooperative air combat method based on deep reinforcement learning. *Acta Automatica Sinica*, 2021, **47**(7): 1610–1623  
(施伟, 冯昞赫, 程光权, 黄红蓝, 黄金才, 刘忠, 等. 基于深度强化学习的多机协同空战方法研究. 自动化学报, 2021, **47**(7): 1610–1623)
  - 62 Liang Xing-Xing, Feng Yang-He, Ma Yang, Cheng Guang-Quan, Huang Jin-Cai, Wang Qi, et al. Deep multi-agent reinforcement learning: A survey. *Acta Automatica Sinica*, 2020, **46**(12): 2537–2557  
(梁星星, 冯昞赫, 马扬, 程光权, 黄金才, 王琦, 等. 多 Agent 深度强化学习综述. 自动化学报, 2020, **46**(12): 2537–2557)
  - 63 Yan D, Weng J, Huang S, Li C, Zhou Y, Su H, et al. Deep reinforcement learning with credit assignment for combinatorial optimization. *Pattern Recognition*, 2022, **124**: Article No. 108466
  - 64 Lansdell B J, Prakash P R, Körding K P. Learning to solve the credit assignment problem. In: Proceedings of the 8th International Conference on Learning Representations. Addis Ababa, Ethiopia: 2020.
  - 65 Sun Chang-Yin, Mu Chao-Xu. Important scientific problems of multi-agent deep reinforcement learning. *Acta Automatica Sinica*, 2020, **46**(7): 1301–1312  
(孙长银, 穆朝絮. 多智能体深度强化学习的若干关键科学问题.



- 自动化学报, 2020, **46**(7): 1301–1312)
- 66 Sunehag P, Lever G, Gruslys A, Czarnecki W M, Zambaldi V, Jaderberg M, et al. Value-decomposition networks for cooperative multi-agent learning based on team reward. In: Proceedings of the 17th International Conference on Autonomous Agents and Multi-agent Systems. Stockholm, Sweden: 2018. 2085–2087
  - 67 Rashid T, Samvelyan M, De Witt C S, Farquhar G, Foerster J N, Whiteson S. QMIX: Monotonic value function factorisation for deep multi-agent reinforcement learning. In: Proceedings of the 35th International Conference on Machine Learning. Stockholm, Sweden: 2018. 4292–4301
  - 68 Son K, Kim D, Kang W J, Hostallero D, Yi Y. QTRAN: Learning to factorize with transformation for cooperative multi-agent reinforcement learning. In: Proceedings of the 36th International Conference on Machine Learning. Long Beach, USA: 2019. 5887–5896
  - 69 Foerster J N, Farquhar G, Afouras T, Nardelli N, Whiteson S. Counterfactual multi-agent policy gradients. In: Proceedings of the 32nd AAAI Conference on Artificial Intelligence. New Orleans, USA: 2018. 2974–2982
  - 70 Nguyen D T, Kumar A, Lau H C. Credit assignment for collective multi-agent RL with global rewards. In: Proceedings of the 32nd International Conference on Neural Information Processing Systems. Montréal, Canada: 2018. 8113–8124
  - 71 Silver D, Hubert T, Schrittwieser J, Antonoglou I, Lai M, Guez A, et al. A general reinforcement learning algorithm that masters chess, Shogi, and go through self-play. *Science*, 2018, **362**(6419): 1140–1144
  - 72 Yu Y. Towards sample efficient reinforcement learning. In: Proceedings of the 27th International Joint Conference on Artificial Intelligence. Stockholm, Sweden: 2018. 5739–5743
  - 73 Ecoffet A, Huizinga J, Lehman J, Stanley K O, Clune J. First return, then explore. *Nature*, 2021, **590**(7847): 580–586
  - 74 Jin C, Krishnamurthy A, Simchowitz M, Yu T C. Reward-free exploration for reinforcement learning. In: Proceedings of the 37th International Conference on Machine Learning. Virtual Event: 2020. 4870–4879
  - 75 Mahajan A, Rashid T, Samvelyan M, Whiteson S. MAVEN: Multi-agent variational exploration. In: Proceedings of the 33rd International Conference on Neural Information Processing Systems. Vancouver, Canada: 2019. Article No. 684
  - 76 Yang Y D, Wen Y, Chen L H, Wang J, Shao K, Mgumi D, et al. Multi-agent determinantal Q-learning. In: Proceedings of the 37th International Conference on Machine Learning. Virtual Event: 2020. 10757–10766
  - 77 Wang T H, Dong H, Lesser V, Zhang C J. ROMA: Role-oriented multi-agent reinforcement learning. In: Proceedings of the 37th International Conference on Machine Learning. Virtual Event: 2020. 9876–9886
  - 78 Zhang Bo, Zhu Jun, Su Hang. Toward the third generation of artificial intelligence. *SCIENTIA SINICA: Informations*, 2020, **50**(9): 1281–1302  
(张钊, 朱军, 苏航. 迈向第三代人工智能. 中国科学: 信息科学, 2020, **50**(9): 1281–1302)
  - 79 Wand Bao-Jian, Hu Da-Sha, Jiang Yu-Ming. Application of improved A\* algorithm in path planning. *Computer Engineering and Applications*, 2021, **57**(12): 243–247  
(王保剑, 胡大猗, 蒋玉明. 改进 A\* 算法在路径规划中的应用. 计算机工程与应用, 2021, **57**(12): 243–247)
  - 80 Zhang Ke, Hao Wen-Ning, Shi Lu-Rong, Yu Xiao-Han, Shao Tian-Hao. Inference of key points of attack in wargame based on cascaded fuzzy system. *Control Engineering of China*, 2021, **28**(7): 1366–1374  
(张可, 郝文宁, 史路荣, 余晓晗, 邵天浩. 基于级联模糊系统的兵棋进攻关键点推理. 控制工程, 2021, **28**(7): 1366–1374)
  - 81 Xing Si-Yuan, Ni Wan-Cheng, Zhang Hai-Dong, Yan Ke. Mining of weapon utility based on the replay data of wargame. *Journal of Command and Control*, 2020, **6**(2): 132–140  
(邢思远, 倪晚成, 张海东, 闫科. 基于兵棋复盘数据的武器效用挖掘. 指挥与控制学报, 2020, **6**(2): 132–140)
  - 82 Jin Zhe-Hao, Liu An-Dong, Yu Li. Hierarchical human-robot cooperative control based on GPR and DRL. *Acta Automatica Sinica*, 2020, **46**: 1–11  
(金哲豪, 刘安东, 俞立. 基于 GPR 和深度强化学习的分层人机协作控制. 自动化学报, 2020, **46**: 1–11)
  - 83 Xu Lei, Yang Yong. Research on evaluation of unit combat action plan based on wargaming. *Fire Control & Command Control*, 2021, **46**(4): 88–92, 98  
(徐磊, 杨勇. 基于兵棋推演的分队战斗行动方案评估. 火力与指挥控制, 2021, **46**(4): 88–92, 98)
  - 84 Li Yun-Long, Zhang Yan-Wei, Wang Zeng-Chen. Technical framework of joint operation scheme deduction and evaluation. *Command Information System and Technology*, 2020, **11**(4): 78–83  
(李云龙, 张艳伟, 王增臣. 联合作战方案推演评估技术框架. 指挥信息系统与技术, 2020, **11**(4): 78–83)
  - 85 Myerson R B. *Game Theory*. Cambridge: Harvard University Press, 2013.
  - 86 Weibull J W. *Evolutionary Game Theory*. Cambridge: MIT Press, 1997.
  - 87 Roughgarden T. Algorithmic game theory. *Communications of the ACM*, 2010, **53**(7): 78–86
  - 88 Chalkiadakis G, Elkind E, Wooldridge M. Cooperative game theory: Basic concepts and computational challenges. *IEEE Intelligent Systems*, 2012, **27**(3): 86–90
  - 89 Zhou Lei, Yin Qi-Yue, Huang Kai-Qi. Game-theoretic learning in human-computer gaming. *Chinese Journal of Computers*, 2022, **45**(9): 1859–1876  
(周雷, 尹奇跃, 黄凯奇. 人机对抗中的博弈学习方法. 计算机学报, 2022, **45**(9): 1859–1876)
  - 90 Lanctot M, Zambaldi V, Gruslys A, Lazaridou A, Tuyls K, Pérolat J, et al. A unified game-theoretic approach to multi-agent reinforcement learning. In: Proceedings of the 31st Conference on Neural Information Processing Systems. Long Beach, USA: 2017. 4190–4203
  - 91 Brown N, Lerer A, Gross S, Sandholm T. Deep counterfactual regret minimization. In: Proceedings of the 36th International Conference on Machine Learning. Long Beach, USA: 2019. 793–802
  - 92 Qiu X P, Sun T X, Xu Y G, Shao Y F, Dai N, Huang X J. Pre-trained models for natural language processing: A survey. *Science China Technological Sciences*, 2020, **63**(10): 1872–1897
  - 93 Zhang Z Y, Han X, Zhou H, Ke P, Gu Y X, Ye D M, et al. CPM: A large-scale generative Chinese pre-trained language model. *AI Open*, 2021, **2**: 93–99
  - 94 Brown T B, Mann B, Ryder N, Subbiah M, Kaplan J, Dhariwal P, et al. Language models are few-shot learners. In: Proceedings of the 34th Conference on Neural Information Processing Systems. Vancouver, Canada: MIT Press, 2020.
  - 95 Meng D Y, Zhao Q, Jiang L. A theoretical understanding of self-paced learning. *Information Sciences*, 2017, **414**: 319–328
  - 96 Singh P, Verma V K, Mazumder P, Carin L, Rai P. Calibrating CNNs for lifelong learning. In: Proceedings of the 34th Conference on Neural Information Processing Systems. Vancouver, Canada: 2020.
  - 97 Cheng W, Yin Q Y, Zhang J G. Opponent strategy recognition in real time strategy game using deep feature fusion neural network. In: Proceedings of the 5th International Conference on Computer and Communication Systems. Shanghai, China: 2020. 134–137
  - 98 Samvelyan M, Rashid T, De Witt C S, Farquhar G, Nardelli N, Rudner T G J, et al. The StarCraft multi-agent challenge. In: Proceedings of the 18th International Conference on Autonomous Agents and Multi-agent Systems. Montreal, Canada: 2019.

2186–2188

- 99 Tang Z T, Shao K, Zhu Y H, Li D, Zhao D B, Huang T W. A review of computational intelligence for StarCraft AI. In: Proceedings of the IEEE Symposium Series on Computational Intelligence. Bangalore, India: 2018. 1167–1173
- 100 Christianos F, Schäfer L, Albrecht S V. Shared experience actor-critic for multi-agent reinforcement learning. In: Proceedings of the Advances in Neural Information Processing Systems 33. Virtual Event: 2020.
- 101 Jaques N, Lazaridou A, Hughes E, Gulcehre C, Ortega P A, Strouse D J, et al. Social influence as intrinsic motivation for multi-agent deep reinforcement learning. In: Proceedings of the 36th International Conference on Machine Learning. Long Beach, USA: 2019. 3040–3049



**尹奇跃** 中国科学院自动化研究所副研究员. 主要研究方向为强化学习, 数据挖掘和人工智能与游戏.

E-mail: qyyin@nlpr.ia.ac.cn

**(YIN Qi-Yue)** Associate professor at the Institute of Automation, Chinese Academy of Sciences. His research

interest covers reinforcement learning, data mining, and artificial intelligence on games.)



**赵美静** 中国科学院自动化研究所副研究员. 主要研究方向为知识表示与建模, 复杂系统决策.

E-mail: meijing.zhao@ia.ac.cn

**(ZHAO Mei-Jing)** Associate professor at the Institute of Automation, Chinese Academy of Sciences. Her

research interest covers knowledge representation and modeling, and complex system decision-making.)



**倪晚成** 中国科学院自动化研究所研究员. 主要研究方向为数据挖掘与知识发现, 复杂系统建模和群体智能博弈决策平台与评估.

E-mail: wancheng.ni@ia.ac.cn

**(NI Wan-Cheng)** Professor at the Institute of Automation, Chinese

Academy of Sciences. Her research interest covers data mining and knowledge discovery, complex system modeling, and swarm intelligence platform and evaluation.)



**张俊格** 中国科学院自动化研究所研究员. 主要研究方向为持续学习, 小样本学习, 博弈决策和强化学习.

E-mail: jgzhang@nlpr.ia.ac.cn

**(ZHANG Jun-Ge)** Professor at the Institute of Automation, Chinese Academy of Sciences. His research

interest covers continuous learning, small sample learning, game decision making, and reinforcement learning.)



**黄凯奇** 中国科学院自动化研究所研究员. 主要研究方向为计算机视觉, 模式识别和认知决策. 本文通信作者.

E-mail: kqhuang@nlpr.ia.ac.cn

**(HUANG Kai-Qi)** Professor at the Institute of Automation, Chinese Academy of Sciences. His research

interest covers computer vision, pattern recognition, and cognitive decision-making. Corresponding author of this paper.)